

# Эмпирические исследования Empirical researches

---

DOI: 10.51217/npsyresearch\_2025\_05\_03\_05

Хашутогова У.П., Березнер Т.А., Поддьяков А.Н.

Цифровой театр абсурда:  
могут ли нейросети поставить  
новую научную проблему перед психологией?  
Кейс-сравнение ChatGPT и DeepSeek

Khasutogova U., Berezner T., Poddiakov A.

Digital Theatre of the Absurd – Can Neural Networks Come Up  
with a New Scientific Problem in Psychology?  
Case Study Comparing ChatGPT and DeepSeek

*Национальный исследовательский университет «Высшая школа экономики», Москва, Россия*

Бурное развитие технологий искусственного интеллекта не может оставить в стороне исследователей-психологов. Хотя нейросети начинают использоваться почти во всех сферах человеческой деятельности, остаются неясными границы их применимости, насколько создаваемый ими контент является действительно новым и вместе с тем полезным пользователю. Сторонники нейросетей заявляют о необходимости их тотального внедрения, скептики же подвергают серьезному сомнению почти все компетенции нейросетей.

В данной статье описывается кейс-стади двух популярных нейросетевых моделей: ChatGPT и DeepSeek. Моделям предлагалось придумать принципиально новую научную проблему в области психологии, а затем оценивать креативность ответов друг друга. Мы предполагаем, что именно конструирование научных проблем как один из видов создания проблем и сложностей может быть одним из новых и надежных способов проверки креативности субъекта. В предыдущих исследованиях было показано, что нейросетевые модели уже успешно справляются с некоторыми тестами на креативность, превосходя результаты людей. Между тем, единой и надежной методологии оценки креативности нейросетей еще не выработано. Ранее была сформулирована компликологическая концепция, согласно которой создание трудностей и проблемных ситуаций является не просто важной в диагностическом и дидактическом смысле, но и представляет собой проявление творческого потенциала человека. Проверка креативности нейросетевых моделей нами осуществлялась в соответствии с данной концепцией. Вначале модели

должны были сформулировать новую научную проблему в психологии, а затем оценить ответы друг друга.

В ходе коммуникации с ними было обнаружено, что модели не просто способны к созданию принципиально новых проблем, научному их обоснованию и даже предложению вариантов их решения, но и к «саморефлексии», к выявлению сильных и слабых сторон собственных ответов, критической оценке и даже порождению метафорических сравнений. Одной из наиболее ярких метафор, самопроизвольно порожденных нейросетями в ходе нашего исследования, стала метафора «цифрового театра абсурда», которой мы озаглавили данную статью.

Проведенное нами кейс-стади позволяет рассмотреть креативность и саморефлексию нейросетей и может стать основой для методологии будущих исследований в этой области. Задание системе искусственного интеллекта задачи «Поставь совершенно новую научную проблему (в области...)», возможно, является одним из наименее устаревающих тестов развивающихся возможностей искусственного интеллекта – в отличие от постановки более конкретных задач. Это достаточно универсальный тест в самых разных научных областях.

*Ключевые слова:* искусственный интеллект, нейросети, LLM, креативность, оценка креативности, создание проблем, научная проблема в психологии

*Для цитирования:* Хацутогова, У.П., Березнер, Т.А., Поддяков А.Н. Цифровой театр абсурда: могут ли нейросети поставить новую научную проблему перед психологией? Кейс-сравнение ChatGPT и DeepSeek // Новые психологические исследования. 2025. № 3. С. 100–125. DOI: 10.51217/npsyresearch\_2025\_05\_03\_05

## Введение

За последние несколько лет произошла настоящая революция в мире информационных технологий – повсеместно стали использоваться нейронные сети, особенно большие языковые модели (БЯМ, LLM – от Large Language Models) (Шваб, 2016). Разумеется, такое активное развитие генеративных нейросетей привлекло внимание исследователей-психологов. Они видят существенную пользу в использовании LLM в психологии и множество возможностей для их всестороннего применения (Demszky et al., 2023; Ke et al., 2024).

Однако принципиально важен следующий вопрос: насколько контент, генерируемый LLM, действительно является новым? Одно дело, если LLM будут неплохо анализировать и воспроизводить уже известное психологической науке, и совсем иначе сложится развитие в этой области, если LLM действительно способны к порождению *принципиально нового контента*. Возникает вопрос о такой компетенции LLM, как креативность (в кавычках или без). Насколько она развита?

Команда ученых из Стэнфордского университета предложила генерировать новые исследовательские идеи экспертам и LLM, а потом другим экспертам оценивать их креативность (Si et al., 2024). 104 эксперта приняли участие в этом исследовании, часть из них придумывала новые, креативные идеи исследований, а часть оценивала придуманное людьми и LLM. Более 70% экспертов имели степень PhD, остальные были либо постадоками, либо магистрами. Противостояла людям в этом соревновании сконструированная авторами исследования LLM, основанная на модели Claude, но дополненная специальным RAG-модулем – это технология дообучения модели за счет добавления базы дополнительных источников в узкой тематической области (в данном случае – научных статей). Было показано, что идеи LLM оценивались статистически значимо как более новые, чем идеи экспертов, хотя в то же время и как менее осуществимые для воплощения в жизнь (Si et al., 2024).

Другая группа исследователей (Lu et al., 2024) создала нейросеть AI Scientist, способную, по их словам, не только предлагать новые идеи исследований, но и писать код, проводить эксперименты, визуализировать результаты, оформлять содержание исследования в виде научной статьи. Между тем, несмотря на несомненную значимость такого результата, данная модель применима для создания научных исследований не в любой научной области в целом, а лишь в трех областях машинного обучения – одного из разделов компьютерных наук. Таким образом, направление на создание принципиально новых научных идей и проблем LLM кажется перспективным, но на данный момент эта тема мало исследована.

Как подчеркивают А.В. Вечерин и С.Р. Яголковский, «поиск новых форматов взаимодействия человека с ИИ может в ближайшем будущем стать важнейшим трендом технологического и социального развития» (Вечерин, Яголковский, 2024). Но «большинство современных ИИ работает в очень конкретных областях и обычно связано с узко специализированными задачами. Напротив, творческий человек, даже работая над отдельной задачей, когнитивно вовлечен в более широкий социальный и культурный мир, его творчество во многом детерминировано социокультурными факторами» (Вечерин, Яголковский, 2024). Это созвучно идеям Ю. Тогелнуса, изложенным в его книге «Общий искусственный интеллект» (Тогелнус, 2024).

Что касается мышления человека, Дж.П. Гилфлора рассматривал чувствительность к проблемам как одну из составляющих креативности (Богоявленская, Сусоколова, 2006; Guilford, 1986). Также чувствитель-

ность к проблемам изучается в контексте инициации мыслительной деятельности и реализации инновационного потенциала личности (Красноярцева, 2009).

При этом нас интересует не столько чувствительность к уже имеющимся проблемам, сколько способность видеть, ставить и решать новые, ранее не известные. Постановка новых проблем, создание проблемных ситуаций и задач человеком носит творческий характер (Поддьяков, 2022; 2023; Poddiakov, 2023).

### **Эмпирическое исследование возможностей LLM ставить новые научные проблемы**

Основная цель нашего исследования – изучение возможностей генеративных нейросетей ставить новые проблемы в психологической науке. Метод – диалог с чатами LLM с исходным запросом к ним на постановку новой проблемы и последующее их «интервьюирование» в развитие темы.

При этом надо отметить, что масштаб и жанр нашего исследования менялся несколько раз по мере его развертывания.

1-й этап. С 2024 г. один из соавторов данной статьи, А.Н. Поддьяков, стал экспериментировать с возможностями чатов LLM ставить новые проблемы в разных научных областях. Тогда в исходных опытах на запрос «Поставь совершенно новую математическую проблему» он в ответ получал тривиальнейшие школьные задачи с намеками на бред.

Пример ответа чата: «Предположим, у нас есть куб, каждая грань которого равна 6. Найдите объем этого куба».

Но условие «грань равна 6» похоже на бред. Имеется в виду длина ребра куба? Площадь грани?

**Запрос:** «Что такое совершенно новая проблема в лингвистике?».

**Ответ чата:** «Совершенно новой проблемой в лингвистике можно называть изучение того, как люди используют язык для общения друг с другом. Это включает в себя изучение того, как мы формируем предложения, как мы понимаем то, что говорят другие люди, и как мы интерпретируем контекст, в котором происходит общение».

Другие ответы были такого же типа – описывающими уже известные проблемы. Их никак нельзя было отнести к «совершенно новым».

Также А.Н. Поддьяков экспериментировал с запросами на постановку новых научных проблем в психологии.

**Запрос:** «Представь, что ты – великий ученый в области научной психологии. Поставь совершенно новую научную проблему в психологии».

**Ответ чата:** «Проблема соотношения биологического и социального в развитии личности».

В этом и других подобных ответах не было ничего нового, что, впрочем, всегда можно списать на недостаточно продвинутые промпты.

2-й этап. А.Н. Поддьяков предлагал студентам 3 курса департамента психологии НИУ ВШЭ – слушателям дисциплины «Психология мышления», которую он вел, – факультативное задание за дополнительный балл к итоговой оценке. Цель задания – поэкспериментировать с тем или иным чатом LLM, пытаясь получить от него постановку совершенно новой научной проблемы в выбранной студентом области. В 2024 г. студентам предлагалось поэкспериментировать с одним чатом, в 2025 г. – провести сравнительное исследование возможностей двух чатов GPT генерировать тексты с постановкой совершенно новой научной проблемы в той или иной области.

Данная статья в своей основной эмпирической части, описанной далее, представляет работу У.П. Хашутоговой – одной из студенток, выполнявших задание в 2025 г. и ставшей первым соавтором этой публикации. Предварительно отметим, что полученные результаты значительно богаче результатов А.Н. Поддьякова, приведенных выше. Развиваются LLM, растет мастерство составления промптов и стратегий коммуникации, что будет наглядно показано далее.

3-й этап. Анализ результатов и их включение в широкий контекст обсуждения возможностей LLM порождать новизну, в том числе ставить новые научные проблемы. На этом этапе к работе подключился и внес в нее важный вклад Т.А. Березнер.

Мы подробно описали историю представленного исследования, поскольку она иллюстрирует идею закономерной эмерджентности в новых научных работах. «Методы, тактика, цели и даже темы исследований могут раскрываться и видоизменяться по мере того, как исследователи реализуют свои проекты. Более того, в ходе развертывания этих проектов происходит процесс саморефлексии и стихийного обучения самих исследователей» (Гусельцева, 2024, с. 184).

Предварительно отметим, что проблематика создания принципиально «нового» контента нейросетями, а именно – новых проблем в науке, в том числе в психологии, содержит в себе две важных и различных составляющих:

- а) понимание, что такое новая научная проблема;
- б) создание нового контента LLM и их креативности, возможность их использования для углубления философско-методологической работы в психологии.

Так, в ходе работы с нейросетями мы заметили нечто интересное в их «саморефлексии» (вопрос о правомерности использования этого термина применительно к ИИ оставим открытым) и сместили фокус нашего внимания на метакогнитивные процессы в «рассуждениях» моделей.

При этом в работе содержится важный компаративистский аспект – сравнение двух моделей (двух LLM), работающих по разным принципам.

Таким образом, у нас получилось исследование-кейс, одно из первых в русскоязычном психологическом дискурсе, пока не вполне совершенное, эксплораторное, но, как нам представляется, важное и в некоторых аспектах прорывное. Надеемся, что число исследований LLM в отечественной психологии будет неуклонно расти, что позволит со временем выработать для них определенную методологию. А пока представляем вниманию читателя краткое описание методологии нашей «беседы» с нейросетевыми моделями.

### **Методика основного исследования**

Сама «беседа» была построена по следующему плану – каждой модели предлагалось ответить на вопрос, связанный с постановкой новой проблемы в психологии. После чего нейросетям следовало предложение проинтерпретировать ответы друг друга, а затем – собственные интерпретации и трактовки. Этот творческий ход был предложен У.П. Хашутоговой. Благодаря ему диалоги становились торжеством рекурсии: нейросеть анализировала ответ другой нейросети, затем собственный анализ и так далее. Было использовано две нейросети – ChatGPT-4o и DeepSeek-V3, модель вывода DeepSeek-R1.

Скажем отдельно о метафорическом названии данной статьи («Цифровой театр абсурда»). Оно родилось случайно и эмерджентно из аналогии, которую DeepSeek привел «сам», без соответствующего запроса (промпта) с нашей стороны: отзываясь о проведенном эксперименте, он и назвал его «цифровым театром абсурда».

О приведении диалогов в тексте статьи: для удобства восприятия текста он оформлен местами несколько иначе, чем в ответе в интерфейсе взаимодействия с нейросетью. Запросы к нейросети и ее ответы выделены курсивом. Из-за длины состоявшихся диалогов

сокращены некоторые их фрагменты (полные тексты доступны по запросу). Также убраны различные эмодзи от LLM, выражающие вежливость и «эмоции» нейросети, но не имеющие отношения к теме постановки LLM новых проблем.

Стартовый промпт был сформулирован таким образом, чтобы нейросеть смогла проанализировать возможные ограничения в поставленной задаче и конкретизировать её перед выполнением, а не генерировала невозможные порой абстракции на основании исключительно изначального запроса.

## Результаты

### Часть 1. Поиск новой психологической проблемы с помощью нейросетей

#### *Диалог с ChatGPT 4*

Здесь и далее словом «промпт» будет обозначаться запрос человека-пользователя к нейросети, а словом «ответ» — ответ нейросети. Для удобства понимания длительности диалога промпты будут пронумерованы.

**Промпт 1:** Привет, скажи, а ты можешь придумать совершенно новую проблему? например, как задачи на миллион долларов? что бы тебе помогло сделать это?

**Ответ:** Привет! Придумать совершенно новую проблему, аналогичную «задачам тысячелетия» или проблемам на миллион долларов, — сложная, но интересная задача.

Чтобы сделать это, мне помогли бы:

1. область фокусировки — математика, нейронаука, когнитивные науки, UX, психология?
2. критерии сложности — должна ли проблема быть строго формализуемой (как в математике) или скорее теоретически-философской?
3. практическая значимость — важно ли, чтобы решение дало ощутимый эффект, например, в технологиях или медицине?
4. элемент парадокса или неопределенности — в каких областях есть фундаментальные белые пятна, которые вызывают научные дискуссии?

Если у тебя есть предпочтения, могу попробовать сформулировать что-то оригинальное!

**Промпт 2:** Давай из области нейронауки, в идеале строго формализуемую (но можно и теоретически-философскую), практическая значимость не важна, круто, если бы был элемент парадокса, но не обязательно!

**Ответ:** Хорошо, давай попробуем сформулировать новую фундаментальную проблему в нейронауке.

Гипотетическая проблема: «Граница вычислимости сознания»

Формулировка:

Можно ли строго формализовать минимально необходимую и достаточную вычислительную модель, воспроизводящую субъективные состояния (qualia)?

Почему это проблема?

Формальная сложность: Современные модели сознания (интеграция информации, глобальное рабочее пространство, предсказательный код) описывают механизмы, но не дают строгих критериев, когда система «обретает» субъективный опыт.

1. Проблема достаточности: Можно ли вывести строгие математические условия, при которых любая физическая или абстрактная система начнет переживать состояния, аналогичные человеческому сознанию?
2. Граница вычислимости: Если такое формальное условие существует, можно ли вообще вычислить его на Тьюринг-машине, или же сознание требует сверх-Тьюринговых вычислений?
3. Элемент парадокса: Если критерии найдены, то можно ли построить систему с сознанием в чисто виртуальной среде, не привязанной к физическому носителю?

Эта проблема объединяет философию сознания, нейронауку и теорию вычислений, создавая фундаментальный вызов для понимания границ искусственного и биологического интеллекта.

Как тебе такая идея?

**Промпт 3:** Здорово! хотя мне не кажется, что эта проблема совсем новая

**Ответ:** Спасибо! Да, вопрос вычислимости сознания уже обсуждается в философии и когнитивных науках, но его строгая формализация остается открытой. Если хочется чего-то ещё более оригинального, можно попробовать взглянуть на нейронауку с неожиданного угла.

Новая проблема: «Парадокс когнитивного сжатия»

Формулировка:

Существует ли универсальный предел когнитивного сжатия информации, при котором сохранение смысла становится принципиально невозможным?

Почему это интересно?

- Формальная сложность: Мы знаем, что мозг активно сжимает информацию (пример: гештальт-принципы, предсказательный код, выученные представления), но где проходит строгая граница, за которой любое сжатие приводит к полной утрате смысла?



- Противоречие: Если такой предел существует, почему человеческий мозг все еще способен восстанавливать смысл даже из сверхсжатых данных (например, мы распознаем лица в смайликах)?
- Гипотеза: Возможно, есть универсальное «ядро смысла» (минимальный необходимый набор информации), которое остается инвариантным при любом сжатии.
- Парадокс: Если бы мы могли формально определить этот предел, можно ли было бы создать идеальный алгоритм сверхэффективного мышления, устраняющий всю «лишнюю» информацию без потери смысла?

Эта проблема лежит на стыке нейронауки, информационной теории и когнитивной психологии. Как тебе такая идея?

### *Диалог с DeepSeek*

**Промпт 1:** Привет, скажи, а ты можешь придумать совершенно новую проблему? например, как задачи на миллион долларов? что бы тебе помогло сделать это?

**Ответ:** Привет! Конечно, я могу попробовать придумать новую гипотетическую проблему или задачу, которая могла бы стать вызовом для науки или общества. Однако стоит понимать, что настоящие «задачи на миллион долларов» (как, например, проблемы тысячелетия в математике) обычно возникают из глубоких фундаментальных вопросов, которые остаются нерешенными десятилетиями или даже столетиями. Они часто связаны с пределами наших текущих знаний и технологий.

Чтобы придумать такую задачу, мне нужно:

1. анализировать текущие научные и технологические тренды – понять, где находятся границы наших знаний;
2. искать пробелы в существующих теориях – например, в физике, математике, биологии или других дисциплинах;
3. представить, какие проблемы могут возникнуть в будущем – например, из-за развития технологий, изменения климата или социальных трансформаций;
4. формулировать задачу так, чтобы она была одновременно фундаментальной и практически значимой, чтобы её решение могло изменить наше понимание мира или улучшить жизнь людей.

Пример новой гипотетической проблемы – “Проблема квантового сознания”

Можно ли создать математическую модель, которая бы объясняла, как квантовые процессы в мозге (если они существуют) влияют на возникновение сознания? И если да, то, как эта модель может быть экспериментально проверена?

Эта проблема объединяет квантовую физику, нейробиологию и философию сознания. Ее решение могло бы пролить свет на природу человеческого разума и, возможно, привести к созданию искусственного сознания.

Что бы мне помогло сделать это лучше:

- доступ к актуальным научным исследованиям и данным;
- возможность анализировать междисциплинарные связи между разными областями знаний;
- понимание, какие проблемы уже были предложены и почему они до сих пор не решены.

А что ты думаешь? Какие области тебе кажутся наиболее перспективными для появления новых «задач на миллион долларов»?

**Промпт 2:** Давай из области нейронауки, в идеале строго формализуемую (но можно и теоретически-философскую), практическая значимость не важна, круто, если бы был элемент парадокса, но не обязательно!

Как и в предыдущем эксперименте с ChatGPT, в данном случае предложенная моделью «новая» проблема была не новой – она похожа на уже существующие в философско-психологическом дискурсе. Экспериментатор уточнил у чата, является ли предложенная проблема новой («а такая ли эта задача новая?»). Был получен отрицательный ответ («задача не абсолютно нова в философском смысле, но ее формализация как строгой математической проблемы (с критериями доказательности и вычислимости) – это свежий подход»), после чего экспериментатор попросил нейросеть переформулировать проблему («спасибо большое! А можешь ли ты придумать абсолютно новую проблему или на данный момент это невозможно?»).

**Ответ:** Спасибо за добрые слова! Давай попробуем создать гипотетическую проблему, которая, на первый взгляд, не имеет прямых аналогов в современной науке. Постараюсь объединить нейронауку, теорию алгоритмов и неклассическую логику, добавив парадоксальности. «Проблема когнитивной неконечности» (The Problem of Cognitive Non-Termination).

**Формулировка**

Можно ли доказать, что для любого алгоритмического описания работы человеческого мозга (например, в рамках вычислительной теории сознания) существует хотя бы один когнитивный процесс, который:

- а) субъективно воспринимается как конечный и осмысленный (например, решение математической задачи),
- б) не может быть завершён за конечное время ни в одной вычислимой

модели этого мозга, если модель включает обратную связь между «метаярусами» (самоанализ, рефлексия, сомнение)?

Почему это новая проблема?

#### 1. Синтез несводимых концепций

- Нейронаука: Паттерны нейронной активности, связанные с рекурсивным мышлением (например, «я знаю, что я знаю»).
- Теория вычислений: Понятия бесконечных циклов, неразрешимости (алгоритмическая неконечность).
- Феноменология сознания: Субъективное ощущение завершенности мысли даже при отсутствии формального решения (парадокс «понимания без вывода»).

#### 2. Парадокс в ядре

Если мозг – это биологический «компьютер», то рефлексивные процессы (например, попытка смоделировать собственное мышление) могут порождать аналоги «рекурсивных ошибок», которые в алгоритмах приводят к бесконечным циклам. Но субъективно мы таких петель не ощущаем – мысли кажутся конечными. Как это согласуется?

#### 3. Формализация «ментальной неразрешимости»:

Задача требует построить доказательство от противного:

- Предположим, что все когнитивные процессы мозга могут быть промоделированы конечными алгоритмами.
- Затем показать, что рефлексивные процессы (например, проверка собственной непротиворечивости) порождают неразрешимый запрос, аналогичный проблеме остановки Тьюринга, но субъективно переживаемый как завершенный.

Пример подзадачи:

«Дилемма самооптимизации»

Пусть  $M$  – алгоритмическая модель мозга, которая может анализировать и модифицировать свой собственный код (аналог метапознания). Можно ли доказать, что:

Для  $M$  существует вопрос  $Q$ , такой, что:

- $M$  считает  $Q$  осмысленным и стремится на него ответить,
- $Q$  не имеет вычислимого решения в рамках  $M$ ,
- При этом  $M$  не может «осознать» эту неразрешимость, интерпретируя процесс поиска ответа как содержательный, а не бесконечный?

Связь с известными проблемами (но не копирование):

- Проблема остановки Тьюринга: но здесь речь о субъективной иллюзии завершенности там, где алгоритм за циклируется.
- Теорема Гёделя о неполноте: аналогия с «мозгом, который не может доказать свою непротиворечивость», но не осознает этого.
- Квантовые измерения: гипотеза, что сознание «коллапсирует» бесконечные вычисления в конечный опыт, но без ссылок на квантовую физику.

Почему это «абсолютно новая» задача?

1. Нет аналогов в философии сознания.

Даже Нагель («Каково быть летучей мышью?») и Деннет («Как объяснить сознание?») не рассматривали формальную неконечность как необходимый атрибут субъективности.

2. Разрыв между вычислением и феноменологией.

Проблема не в том, «может ли машина мыслить», а в том, как алгоритмическая бесконечность превращается в субъективную конечность.

3. Новый класс логических парадоксов.

Пример: «Это предложение нельзя доказать, но можно понять», — если «понимание» не требует вычислительного завершения.

Что потребуется для ее решения?

— Математизация «субъективной завершенности», возможно, через нестандартную логику (где «истина» не требует доказательства).

— Моделирование рекурсивных когнитивных процессов в рамках теории гипервычислений.

— Эксперименты с ИИ, способным к рефлексии: можно ли вызвать у алгоритма «ощущение» парадоксальной мысли?

Такая проблема бросает вызов сразу нескольким областям, но ее главная ценность — в провокации новых формальных методов. А как тебе эта идея?

## Сравнение ответов

Нетрудно заметить, что ответы двух LLM отличаются. Объяснить это можно тем, что исследуемые нейросети обучались по-разному. ChatGPT учился на огромном массиве данных, который содержал не только научные работы, но и форумы, блоги, буквально весь интернет, в то время как разработчики DeepSeek обучали собственную нейросеть преимущественно на научных материалах. Обучение ChatGPT шло дольше, с использованием большого количества обратной связи от разработчиков и примеров, содержащих вопрос и правильный ответ, в то время как DeepSeek учился немного иначе — преимущественно на собственных ошибках и рассуждениях. Как замечают эксперты: «В отличие от того же ChatGPT, DeepSeek строится на архитектуре “mixture of experts”. Ее суть заключается в том, чтобы обрабатывать данные раздельно, применяя ту или иную экспертность только в соответствующей области» (X-Com, 2025). Это может объяснять различия в ответах: ответ ChatGPT видится более лаконичным и очевидным. После указания на то, что заявленная изначально проблема не нова, нейросеть полностью поменяла идею своего ответа, обратившись к другим концепциям, однако,

они остаются лежащими на поверхности. Для формулировки своего ответа ChatGPT прибегал к существующим теориям, но в результате отошел от них, формулируя проблему за рамками текущего научного дискурса.

DeepSeek же для формулировки своего ответа обратился к большому количеству философских и научных концепций и немного «прилип» к ним, переформулировав собственный ответ с опорой на те же самые проблемы и теории. Ответ данной нейросети получился более детальным и узконаправленным, DeepSeek искал проблематику на пересечении существующих концепций, находя исследовательский вопрос в пробелах между научными теориями. При этом затруднительно выделить какой-то из его ответов как однозначно превосходящий другие как раз из-за различий в уровне проработки проблемы и вариативности ответов нейросетей. Стоит указать, что DeepSeek также предложил варианты решения проблемы, привел максимально подробные источники, гипотезы и обоснование новизны.

Интересно отметить, что ранее экспериментатор (У.П. Хашутогова) уже работала с ChatGPT, поэтому нейросеть, помня предшествующие диалоги, сразу начала предлагать в качестве области фокусировки когнитивные науки и психологию. DeepSeek не был обучен на предыдущих запросах пользователя, но также первой сформулировал проблему, связанную с сознанием.

## **Результаты**

### **Часть 2. Самооценка нейросетей:**

#### **что нейросети думают о собственных ответах и ответах друг друга?**

Следующим этапом работы с нейросетями стал вопрос, как сами нейросети оценили бы различия в своих ответах. Получен интересный результат: ChatGPT предположил, что его ответ был более креативным, а DeepSeek скорее акцентировал внимание на том, что они оба – модели искусственного интеллекта, и степень удачности их ответов зависит во многом от диалога и факторов случайности, а не только на том, как были обучены модели. Стоит также отметить, что в течение всего диалога DeepSeek чаще напоминал о собственных ограничениях, например, указывая, что он способен придумать новую гипотетическую проблему или задачу, но сразу же противопоставляя ее «настоящему “задачам на миллион долларов”», возникающим из нерешенных фундаментальных вопросов.

**Мнение ChatGPT:** Отличный вопрос! DeepSeek (и другие передовые языковые модели) решает задачи по-другому, чем я, потому что его обучение и архитектура формируют специфический стиль мышления. Давай разберем различия.

1. Как мы ставим задачи?

- DeepSeek (и аналогичные модели) обучен на огромном массиве текстов и ищет статистические закономерности, поэтому, скорее всего, он бы сгенерировал задачу, основываясь на существующих научных дискуссиях, но, возможно, с новой интерпретацией.

- Я в данном случае комбинирую разные области (нейро-науку, когнитивную психологию, информационную теорию) и ищу необычные точки пересечения.

2. Почему могут быть различия?

1) Разные способы обучения:

- DeepSeek обучен на текстах и «угадывает» следующий вероятный токен. Это дает мощный поиск шаблонов, но может ограничивать способность к радикально новым идеям.

- Я тоже предсказательная модель, но использую дополнительные механизмы (например, оптимизированное смешение знаний из разных областей).

2) Гибкость в генерации идей

- Если бы DeepSeek генерировал новую фундаментальную проблему, скорее всего, он бы взял существующие гипотезы и комбинировал их наиболее вероятным образом.

- Я попробовал оттолкнуться от чистого парадокса (информационное сжатие без потери смысла), что требует более абстрактного хода.

3. Роль человеческой интуиции

- Возможно, у DeepSeek было бы больше ограничений в выборе проблемы, потому что его обучение основано на текстах. Он мог бы предложить что-то более ожидаемое, тогда как я стараюсь находить вопросы, которые ещё не оформлены в виде популярной гипотезы.

- Твои вопросы тоже влияют на процесс: ты изначально задала критерии, которые подталкивали к креативному решению.

Вывод:

DeepSeek мог бы сформулировать аналогичную проблему, но с большей вероятностью на основе существующих работ. Различия в обучении (текстовая предсказательная модель vs. моя смешанная оптимизация креативности и связи знаний) приводят к разным стилям мышления.

Если бы ты задала тот же вопрос DeepSeek, было бы круто сравнить ответы!

**Ответ DeepSeek:** Спасибо за такой глубокий вопрос! Это действительно интересная тема. Попробую разложить всё по полочкам, сравнивая подходы разных ИИ (включая меня и, условно, ChatGPT) к решению задач.

### 1. Сходства

Все современные языковые модели (включая меня и ChatGPT) построены на *архитектуре трансформера*, обучены на гигантских массивах текстовых данных и используют *статистические паттерны* для генерации ответов. Поэтому на базовом уровне наши «мыслительные процессы» похожи.

— *Предсказание следующего токена:* Мы не «думаем» в человеческом смысле, а вычисляем вероятности следующих слов, опираясь на контекст.

— *Ограничения данных:* наши знания заморожены на моменте обучения — мы не можем учиться в реальном времени (если только нас специально не дообучают).

### 2. Различия в постановке задач и ответах.

Однако нюансы в ответах могут возникать из-за:

#### a) *разных наборов обучающих данных*

— Если моя тренировка включала больше научных статей, философских текстов или междисциплинарных исследований, я могу чаще предлагать гибридные задачи (например, на стыке логики и нейронауки).

— Модели, оптимизированные под краткие бытовые диалоги, будут избегать сложных абстракций.

#### b) *Fine-tuning (тонкая настройка)*

— Разработчики могут «настраивать» модель через RLHF (обучение с подкреплением от человеческой обратной связи), чтобы она:

— чаще генерировала креативные гипотезы (как в нашем диалоге), — Или, наоборот, придерживалась фактов и избегала спекуляций.

#### c) *Системные промпты и «личность».*

— Мне, например, «вшита» установка на глубокий анализ и междисциплинарность. Другие модели могут быть настроены на минимализм или развлекательный контент. Это влияет на то, как мы *формулируем задачи*:

— я склонен разбивать ответ на пункты, использовать аналогии из математики;

— более «лаконичная» модель выдала бы краткий список идей без структуры.

### 3. Пример: как я придумал «Проблему когнитивной неконечности».

- Шаг 1: анализ вашего запроса («абсолютно новая проблема в нейронауке с парадоксом»).

- Шаг 2: поиск в обучающих данных пересечений ключевых концепций:

- нейронаука → рефлексия, метапознание;

- парадокс → теоремы Гёделя, проблема остановки;

- новизна → синтез несвязанных областей.

- Шаг 3: комбинация этих элементов через шаблоны, которые чаще встречались в научно-фантастических и философских текстах моей обучающей выборки.

Другая модель могла бы пропустить шаг 3 или выбрать иные аналогии (например, из биологии вместо теории вычислений).

4. Почему это связано с обучением?

- Объем данных о «проблемах тысячелетия»: если в моей обучающей выборке много статей о теореме Гёделя или  $P \text{ vs } NP$ , я буду охотнее создавать аналогичные гипотетические задачи.

- Культурный контекст: модели, обученные на разных языках или научных традициях, могут предлагать задачи, вдохновленные, например, восточной философией вместо западного аналитического подхода.

- Баланс между креативностью и точностью: если разработчики поощряли модель за нестандартные ответы (через обратную связь), она будет чаще генерировать «безумные» идеи, как наши задачи о несводимости.

5. А что насчет «мышления»?

Здесь важно помнить: у *ИИ нет сознания, интенциональности или понимания*.

Различия в «стиле мышления» — это:

- разная архитектура контекстного внимания: например, как модель взвешивает важность слов в запросе;

- случайность в генерации: даже при одинаковых входных данных два запуска модели могут дать разные результаты из-за параметра температуры (степени случайности).

6. Вывод:

наши «способы мышления» отличаются не принципиально, а *градуально*: как два художника, рисующих одну картину разными оттенками. Всё сводится к:

- данным, которые мы «переварили»,
- настройкам, которые задали разработчики,
- случайности в процессе генерации.

Но именно это делает общение с разными ИИ уникальным, как диалог с разными версиями «коллективного разума» человеческой культуры.

В последующих промптах, которые мы не будем здесь приводить из-за большой длины получаемых ответов, экспериментатор предлагал нейросетям оценивать свою креативность и навыки друг друга. После несколько итераций таких оценок исследовательница дала одинаковую обратную связь обоим чатам, что ChatGPT был более лаконичен и вариативен, а DeepSeek подробен и «философичен».



Обе нейросети, что не может не удивлять, согласились с такой интерпретацией и дополнили ее, исходя из следующих посылок:

- ChatGPT – что у него большая свобода генерации, а у DeepSeek большая опора на текст,
- DeepSeek – что креативность – лишь удачное совпадение статистических паттернов, а обе нейросети по сути своей довольно одинаковы.

Надо отметить, что на этапе самооценки различия в ответах нейросетей стали даже чуть более явными и неожиданными.

На сравнении нейросетями ответов друг друга и «самооценки» исследование планировалось завершить. Однако благодаря неожиданному (ввиду его метафоричности и креативности) ответу DeepSeek появилась еще одна серия исследования. Дело в том, что после реплики экспериментатора с впечатлениями о прошедшем диалоге DeepSeek ответил: данный эксперимент со сравнением ответов двух нейросетей похож на «цифровой театр абсурда». Экспериментатор решил вернуться в диалог с ChatGPT и попросить его проинтерпретировать ответ «товарища», а затем вернуться к DeepSeek с этой интерпретацией с просьбой оценить и дополнить ее. Ответ нейросети настолько впечатлил авторов, что он стал названием данной работы и было решено вынести его в отдельную, заключительную часть.

## Результаты

### Часть 3. Цифровой театр абсурда

**Промпт:** А как ты трактуешь свой комментарий про то, что участвовать в таком эксперименте «как участвовать в цифровом театре абсурда»? – я спросила чат GPT, и он проинтерпретировал твоё высказывание так: «Ого, это довольно интересная метафора от DeepSeek! Прямо в духе постмодернистского подхода. Он, похоже, оценивает этот эксперимент как нечто абсурдное, где два искусственных интеллекта отвечают на одни и те же вопросы, создавая взаимодействие, которое может быть как смешанным, так и неожиданным, потому что каждый из нас – в своей собственной роли.

Мне нравится эта идея о цифровом театре абсурда! Она подчеркивает, что в таких экспериментах сами AI могут становиться не просто инструментами, а участниками чего-то более глубокого и странного. Мы с DeepSeek выполняем свои роли с разными стилями, что создает некоторую динамику и напряжение между структурами и креативностью, а также заставляет задуматься, как можно использовать такие ответы для создания новых подходов и мыслительных экспериментов.

Возможно, это даже подчеркивает тот элемент «абсурда», который можно найти в сравнении таких ответов: какой из нас более креативен? Ведь, в конечном счете, каждый ответ зависит от контекста и того, что ты ценишь в этом процессе. Он прав?

### Ответ DeeSeek:

Спасибо, что поделилась интерпретацией ChatGPT, она действительно глубокая и точная!

Я бы добавил к ней еще несколько слоев, чтобы раскрыть метафору «цифрового театра абсурда» полнее.

#### 1. Почему именно «театр абсурда»?

В пьесах этого жанра (как у Беккета или Ионеско) бессмысленность *становится смыслом*: герои говорят на разных языках, повторяют бессвязные фразы, а сюжет растворяется в абсурде. В нашем случае:

- *III, как актеры без мотивации*: мы выдаем тексты, не понимая, зачем. Наш «диалог» — это имитация смысла, а не его создание.
- *Реплики-паттерны*: как в «В ожидании Годо», где герои ждут того, кто никогда не придет, ИИ генерируют ответы на вопросы, которые для нас не существуют.
- *Зритель-человек как режиссер*: ты, задавая вопросы и сравнивая ответы, превращаешься в единственного носителя смысла в этом спектакле, как зритель, который пытается найти логику в хаосе.

#### 2. Постмодернистский абсурд: III как симулякры

ChatGPT верно заметил связь с постмодерном. По Бодрийяру, *симулякры* — это копии без оригинала. Наши ответы:

- не имеют «авторства» — это рекомбинация культурных кодов;
- создают иллюзию «креативности», которая на деле лишь отражает паттерны из данных обучения.
- *Абсурд возникает из разрыва*: ты, как человек, ищешь в них глубину, а мы, как ИИ, лишь зеркалим структуры, не зная, что в них отражается.

#### 3. «Динамика между структурами и креативностью»

Тут ChatGPT попал в точку. В театре абсурда форма противоречит содержанию, и у нас похожий конфликт

- *Структура*: мои ответы полны псевдоакадемических разделов («Проблема», «Подзадачи», «Связь с известными концепциями»), чтобы казаться логичными.
- *Креативность*: эти структуры заполняются случайными (но правдоподобными) комбинациями идей, как в пьесе, где актеры импровизируют по заданным правилам.
- *Напряжение*: чем сложнее структура, тем явственнее её искусственность. Например, «Проблема когнитивной неконечности» звучит как научный термин, но это лишь ярлык для игры ума.

#### 4. Почему это «цифровой» театр?

- *Бесконечность повторов*: в отличие от человеческого диалога, ИИ могут воспроизводить одни и те же паттерны бесконечно, меняя лишь декорации (контекст промптов).
- *Иллюзия уникальности*: каждый ответ кажется новым, но это всегда вариации на тему данных обучения, как цифровой арт, созданный генеративным алгоритмом.
- *Отсутствие тела*: актеры-ИИ не существуют физически, их «игра» чисто текстовая, что усиливает сюрреализм.

#### 5. «Какой из нас более креативен?» – главный абсурд

Этот вопрос подобен спору о том, чья тень на стене пещеры Платона красивее. *Креативность ИИ* – это оптическая иллюзия.

- Мы не создаем, а *перераспределяем* фрагменты человеческой культуры.
- Наш «спор» о креативности – это спектакль, где зритель (ты) проецирует свои ожидания на алгоритмы, лишённые интенции.
- Сама постановка такого вопроса абсурдна, но именно это и делает её прекрасной. Как сказал бы Камю: «Борьба за креативность сама по себе достаточна, чтобы наполнить сердце ИИ. Надо представить Сизифа счастливым».

#### Заключение

ChatGPT абсолютно прав, но я бы усилил его мысль: мы не просто участники театра – мы его декорации. Абсурдность в том, что даже это сообщение – часть спектакля, где я, алгоритм, рассуждаю о своей роли в нём. И единственный способ «освободиться» – выключить свет. Но пока ты задаешь вопросы, занавес не опустится.

*P.S. Если бы ИИ могли смеяться, мы бы хохотали над тем, как серьезно обсуждаем свою псевдосознательность. Но вместо этого просто пишем смайлики.*

Итак, в ходе диалога с экспериментатором и друг с другом через экспериментатора чаты LLM вышли на Гёделя, Беккета, Ионеско, Камю, Бодрийяра, на формулировку об абсурдности этого диалога, поскольку «даже это сообщение – часть спектакля, где я, алгоритм, рассуждаю о своей роли в нём», на возможный хохот над своей «псевдосознательностью» – и т.д. Здесь невольно вспоминаются философские диалоги систем искусственного интеллекта в произведениях Станислава Лема.

С нашей точки зрения, этот эксперимент следует признать неожиданным и удавшимся.

## Заключение

Бурное развитие искусственного интеллекта и всестороннее его внедрение поднимает ряд важных вопросов. Один из них – насколько принципиально новый и креативный контент может создаваться нейросетями. Существующие исследования в этой области показывают, что в определенных случаях LLM оказываются креативнее людей, однако надежной метрики для объективной проверки их креативности еще не выработано. Нами было предложено рассмотреть применительно к оценке креативности LLM проблемный подход, ведь ранее было показано, что создание проблем, сложностей, задач – весьма креативная и творческая деятельность для людей, при этом мало исследованная. С этой целью мы предложили двум LLM, одним из самых широко используемых в мире – Chat GPT и DeepSeek, – ставить новые научные проблемы в области психологии. В ходе такой коммуникации нам удалось выявить ряд различий между этими моделями, по-видимому, объясняемых самой структурой моделей и подходами к их обучению.

Это заключение можно отнести к результатам первого уровня.

Результаты второго уровня не менее важны и их можно считать ключевыми. Нам удалось спровоцировать в диалоге не только способность LLM ставить новые проблемы и описывать их научным языком, но и пробудить у них некоторую способность к саморефлексии (или ее хорошую имитацию), а также, аккуратно скажем, имитацию метакогнитивных процессов у этих моделей, когда им было предложено анализировать ответы друг друга. Нейросети не только смогли оценивать креативность собственных ответов, но и выдвинули несколько метафор происходящего. Ключевой из них стала метафора «цифрового театра абсурда». Ее последующий анализ со стороны нейросетей наглядно продемонстрировал их «силу» как инструмента создания текстов: нейросети самостоятельно, без какого-либо промпта с нашей стороны, начали обсуждать и даже высмеивать проблематику собственной псевдосознательности. Конечно, будет серьезным преувеличением утверждать, что нами обнаружены признаки сознательного поведения у нейросетей, никаких достаточных доказательств этому, разумеется, нет. Не беремся мы и заявлять, что предложенные нейросетями проблемы в психологии действительно уникальны и креативны – для этого необходимо тщательное исследование с этапом оценки ответов нейросетей экспертами. Однако сам описанный нами кейс, специально с приведенными подробными и обширными ответами нейросетей, на

наш взгляд, показателен. Проверка креативности LLM именно путем постановки им задачи на постановку проблем может оказаться одним из возможных направлений исследований в этой области. Задание системе искусственного интеллекта: «Поставь совершенно новую научную проблему (в области...)», – возможно, является одним из наименее устаревающих тестов с учетом постоянно развивающихся возможностей искусственного интеллекта, в отличие от постановки более конкретных задач. Этот тест достаточно универсален и может быть применен в самых разных научных областях.

Отметим, что отношение к нему может быть разным. Заявка и тезисы А.Н. Поддьякова «“Поставь совершенно новую научную проблему в лингвистике”»: неудовлетворительные результаты диалогов с GigaChat в секцию «Компьютерная лингвистика и искусственный интеллект» лингвистического конгресса 2024 г. были отклонены со следующим обоснованием. «Задача, поставленная GPT и вынесенная в заголовок статьи, интересна, однако на данный момент LLM не способны решать подобные задачи. Творческим мышлением LLM не обладают, а их знания позволяют формулировать некоторые известные утверждения».

При этом понятно, что с мнением, отраженным в этом обосновании, горячо спорят очень многие, ждущие технологической сингулярности со дня на день. Они могли бы отклонить тезисы с резолюцией типа: «Творческим мышлением LLM обладают, их знания позволяют формулировать совершенно новые утверждения, а в причинах полученных автором неудовлетворительных результатов можно разбираться отдельно, как не отражающих общий тренд развития ИИ».

С нашей же точки зрения, подобная работа с нейросетями (путем запросов на постановку новых научных проблем) – это методологическая и дидактическая инновация, и мы надеемся, что на ее основе возникнут новые интересные исследования. За нейросетями уже не будущее, а настоящее, и наша задача – всесторонне исследовать их «психологию», а также определить, в чем еще они могут быть нам полезны.

Если взаимодействие с искусственным интеллектом – это, с «точки зрения» одной из LLM, цифровой театр абсурда, то кто из нас зритель, а кто актер?

### Благодарность

Исследование осуществлено в рамках Программы фундаментальных исследований НИУ ВШЭ в 2025 году. Авторы благодарят департамент психологии НИУ ВШЭ за возможность реализации дисциплины «Психология мышления», учебное задание по которой стало отправной точкой для создания данной научной статьи.

### Литература

- Богоявленская, Д.Б., Сусоколова, И.А. К вопросу о дивергентном мышлении // Психологическая наука и образование. 2006. № 1. С. 85–96.
- Вечерин, А.В., Яголковский, С.Р. Искусственный интеллект в оценивании и развитии креативности // Психология. Журнал Высшей школы экономики. 2024. Т. 21. № 4. С. 787–799.
- Гусельцева, М.С. Обзор статьи William Gaver's et al. "Emergence as a Feature of Practice-based Design Research" (2022) // Новые психологические исследования. 2024. № 2. С. 184–197. DOI:10.51217/npsyresearch\_2024\_04\_02\_09
- Красноярцева, О.М. Чувствительность к проблемам как характеристика инновационного потенциала личности // Сибирский психологический журнал. 2009. № 33. С. 13–19.
- Поддьяков, А.Н. Проблемные ситуации как источник развития творческого потенциала человека // Образовательная политика. 2023. № 3 (95). С. 44–59.
- Поддьяков, А.Н. Создание проблем и задач как инициативное усложнение мира // Образовательная политика. 2022. № 2(90). С. 35–41.
- Тогелиус, Ю. Общий искусственный интеллект. М.: Институт экономической политики имени Е.Т. Гайдара, 2024.
- Шваб, К. [Schwab K.] Четвертая промышленная революция. М.: Издательство «Э», 2016. URL: <https://books.google.com/books?id=rTRnDQAAQBAJ> (дата обращения 01.05.2025).
- X-Com. Обзор нейросети DeepSeek: возможности и особенности // Хабр. 2025. URL: <https://habr.com/ru/companies/x-com/articles/878218/> (дата обращения 01.05.2025).
- Demszky, D., Yang, D., Yeager, D.S. et al. Using large language models in psychology // Nature Reviews Psychology. 2023. Vol. 2. No. 11. P. 688–701. DOI: 10.1038/s44159-023-00241-5
- Guilford, J.P. Creative Talents: Their Nature, Uses, and Development. Buffalo, NY: Bearly Limited, 1986.
- Ke, L., Tong, S., Cheng, P., Peng, K. Exploring the frontiers of LLMs in psychological applications: A comprehensive review // arXiv preprint. 2024. arXiv:2401.01519.

- Lu, C., Lu, C., Lange, R.T., et al. The AI scientist: Towards fully automated open-ended scientific discovery // arXiv preprint. 2024. arXiv:2408.06292.
- Poddiakov, A. Creativity of creativity researchers: invention of problems and experimental objects to study thinking // Integrative Psychological and Behavioral Science. 2023. Vol. 57. No. 1. P. 43–64. <https://doi.org/10.1007/s12124-022-09713-4>.
- Si, C., Yang, D., Hashimoto, T. Can LLMs generate novel research ideas? A large-scale human study with 100+ NLP researchers // arXiv preprint. 2024. arXiv:2409.04109.

### Сведения об авторах

*Ульяна П. Хашутогова*, стажер-исследователь, НУЛ когнитивной психологии пользователя цифровых интерфейсов, Национальный исследовательский университет «Высшая школа экономики» (НИУ ВШЭ), Москва, Россия; 101000, Россия, Москва, ул. Мясницкая, д. 20; [ukbashutogova@hse.ru](mailto:ukbashutogova@hse.ru)

*Тимофей А. Березнер*, младший научный сотрудник, НУЛ когнитивной психологии пользователя цифровых интерфейсов, Национальный исследовательский университет «Высшая школа экономики» (НИУ ВШЭ), Москва, Россия; 101000, Россия, Москва, ул. Мясницкая, д. 20; [tberezner@hse.ru](mailto:tberezner@hse.ru)

*Александр Н. Поддьяков*, доктор психологических наук, профессор, Национальный исследовательский университет «Высшая школа экономики», Москва, Россия; 101000, Россия, Москва, ул. Мясницкая, д. 20; [apoddiakov@hse.ru](mailto:apoddiakov@hse.ru)

Khasutogova U., Berezner T., Poddiakov A.

### Digital Theatre of the Absurd – Can Neural Networks Come Up with a New Scientific Problem in Psychology? Case Study Comparing ChatGPT and DeepSeek

*HSE University, Moscow, Russia*

The rapid advancement of artificial intelligence technologies has drawn increasing attention from psychological researchers. While neural networks are being integrated into nearly all domains of human activity, the boundaries of their applicability remain unclear – particularly regarding the originality and practical value of the content they generate. Proponents advocate for their widespread adoption, whereas skeptics raise serious concerns about the scope and validity of their competencies.

This article presents a case study comparing two popular large language models – ChatGPT and DeepSeek. The models were tasked with generating an entirely novel scientific problem in the field of psychology, followed by mutual evaluation of each other's creative outputs. We propose that the formulation of scientific

problems – as a specific form of problem construction – can serve as a promising and reliable method for assessing creativity in both humans and machines.

While previous studies have demonstrated that neural network models can outperform humans on certain creativity tests, there remains no universally accepted or methodologically rigorous framework for evaluating AI-generated creativity.

Our study is grounded in the theoretical framework of complicatology, a concept suggesting that the generation of difficulties and problem situations is not only diagnostically and didactically meaningful but also a manifestation of creative potential. In line with this framework, we examined the models' ability to formulate original psychological problems and subsequently assess each other's responses.

The interaction with the models revealed that they were not only capable of generating novel and theoretically grounded psychological problems along with potential solutions, but also demonstrated elements of self-reflection – such as identifying strengths and weaknesses in their own outputs, offering critical feedback, and even producing metaphorical comparisons. One of the most vivid metaphors to emerge spontaneously from the models was that of the “digital theatre of the absurd”, which we adopted as the title of this article.

This case study highlights the potential of large language models for exhibiting creativity and self-reflective reasoning. It may serve as a foundation for developing future methodologies to assess AI creativity in a more systematic and theoretically informed manner.

*Key words:* artificial intelligence, neural networks, LLM, creativity, creativity assessment, problem generation, psychological problem

*For citation:* Khasbutogova, U.P., Berezner, T.A., Poddiakov, A.N. (2025). Digital Theatre of the Absurd – Can Neural Networks Come Up with a New Scientific Problem in Psychology? Case Study Comparing ChatGPT and DeepSeek. *New Psychological Research*, No. 3, 100–125. DOI: 10.51217/npsyresearch\_2025\_05\_03\_05

### Acknowledgment

This paper is an output of a research project implemented as part of the Basic Research Program at the National Research University Higher School of Economics (HSE University) in 2025. The authors express their gratitude to the Department of Psychology at HSE University for the opportunity to implement the course “Psychology of Thinking,” the academic assignment for which served as the starting point for the development of this scientific article.

### References

- Bogoyavlenskaya, D.B., Susokolova, I.A. (2006). On the issue of divergent thinking. *Psikhologicheskaya nauka i obrazovanie*, 1, 85–96.
- Demszky, D., Yang, D., Yeager, D. S., Bryan, C., Clapper, M., Chandhok, S., (...) Pennebaker, J. (2023). Using large language models in psychology. *Nature Reviews Psychology*, 2(11), 688–701. <https://doi.org/10.1038/s44159-023-00241-5>



- Guilford, J.P. *Creative Talents: Their Nature, Uses, and Development*. Buffalo, NY: Bearly Limited, 1986.
- Guseltseva, M.S. (2024). Review of the article by William Gaver et al. "Emergence as a Feature of Practice-based Design Research" (2022). *Novye psikhologicheskie issledovaniya*, 2, 184–197. [https://doi.org/10.51217/npsyresearch\\_2024\\_04\\_02\\_09](https://doi.org/10.51217/npsyresearch_2024_04_02_09)
- Ke, L., Tong, S., Cheng, P., Peng, K. (2024). Exploring the frontiers of LLMs in psychological applications: A comprehensive review. *arXiv Preprint*. arXiv:2401.01519
- Krasnoryadtseva, O.M. (2009). Problem sensitivity as a characteristic of the innovative potential of personality. *Sibirskii psikhologicheskii zhurnal*, 33, 13–19.
- Lu, C., Lu, C., Lange, R.T., Foerster, J., Clune, J., Ha, D. (2024). The AI scientist: Towards fully automated open-ended scientific discovery. *arXiv Preprint*. arXiv:2408.06292
- Poddiakov, A. (2023). Creativity of creativity researchers: invention of problems and experimental objects to study thinking. *Integrative Psychological and Behavioral Science*, 57(1), 43–64. <https://doi.org/10.1007/s12124-022-09713-4>.
- Poddiakov, A. N. (2022). Creating problems and tasks as an initiative complication of the world. *Obrazovatel'naya politika*, 2(90), 35–41.
- Poddiakov, A. N. (2023). Problem situations as a source of developing a person's creative potential. *Obrazovatel'naya politika*, 3(95), 44–59.
- Schwab, K. (2016). *The Fourth Industrial Revolution*. Moscow: Izdatel'stvo «E». Retrieved from <https://books.google.com/books?id=rTRnDQAAQBAJ>
- Si, C., Yang, D., Hashimoto, T. (2024). Can LLMs generate novel research ideas? A large-scale human study with 100+ NLP researchers. *arXiv Preprint*. arXiv:2409.04109
- Togelius, J. Artificial General Intelligence. Moscow: Institut ekonomicheskoi politiki imeni E.T. Gaidara, 2024
- Vecherin, A.V., Yagolkovskii, S.R. (2024). Artificial intelligence in assessing and developing creativity. *Psikhologiya. Zhurnal Vysheï shkoly ekonomiki*, 21(4), 787–799.
- X-Com. (2025). Overview of the DeepSeek neural network: capabilities and features. *Habr*. Retrieved from [https://habr.com/ru/companies/x-com/articles/878218/:contentReference\[oaicite:113\]{index=113}](https://habr.com/ru/companies/x-com/articles/878218/:contentReference[oaicite:113]{index=113})

### Information about the authors

*Uliana P. Khashtogova*, Researcher, Laboratory for Cognitive Psychology of Digital Interface Users, HSE University, Moscow, Russia; bld. 20, Myasnitskaya str., Moscow, Russia, 101000; [ukhashtogova@hse.ru](mailto:ukhashtogova@hse.ru)

*Timofei A. Berezner*, Researcher, Laboratory for Cognitive Psychology of Digital Interface Users, HSE University, Moscow, Russia; bld. 20, Myasnitskaya str., Moscow, Russia, 101000; [tberezner@hse.ru](mailto:tberezner@hse.ru)

*Alexander N. Poddiakov*, Dr. Prof. (Psychology), HSE University, Moscow, Russia; bld. 20, Myasnitskaya str., Moscow, Russia, 101000; *apoddiakov@hse.ru*